

Article

Bridge Acceleration Data Cleaning Based on Two-Stage Classification Model with Multiple Feature Fusion

Yichao Xu ^{1,2,3}, Yufeng Zhang ^{2,3} and Jian Zhang ^{1,4,*} ¹ School of Civil Engineering, Southeast University, Nanjing 210096, China; xyc576@jsti.com² Jiangsu Transportation Institute Group, Nanjing 211112, China³ National Key Laboratory of Safety, Durability and Healthy Operation of Long Span Bridges, Nanjing 211112, China⁴ Jiangsu Key Laboratory of Engineering Mechanics, Southeast University, Nanjing 210096, China

* Correspondence: jian@seu.edu.cn

Abstract: Over the past few decades, rapid economic development has led to the establishment of numerous monitoring systems, resulting in the accumulation of vast amounts of monitoring data. Among these data, dynamic acceleration data stand out prominently. However, the quality of collected acceleration data is often compromised due to factors such as challenging operational environments and sensor malfunctions. This severely hampers the value extracted from the data. Although manual identification and classification of data anomalies are more reliable, they are time consuming and labor intensive. To address the challenge of identifying and classifying anomalies in massive acceleration data, this paper proposes a two-stage model for intelligent data cleaning. Firstly, raw acceleration data are transformed into IPDF and PSD features, and a one-dimensional convolutional neural network is trained to preliminarily identify and classify acceleration data anomalies. Subsequently, the RPV indicator is extracted from the original data of the normal and outlier categories to achieve precise classification based on threshold values. The proposed method is successfully validated using acceleration monitoring data from a large-span arch bridge, achieving an accuracy of over 99%. Furthermore, compared to directly employing a one-dimensional CNN classification model, the approach significantly enhances the model's perception of local significant disturbances.



Citation: Xu, Y.; Zhang, Y.; Zhang, J. Bridge Acceleration Data Cleaning Based on Two-Stage Classification Model with Multiple Feature Fusion. *Appl. Sci.* **2023**, *13*, 12045. <https://doi.org/10.3390/app132112045>

Academic Editor: Carmelo Gentile

Received: 9 September 2023

Revised: 28 October 2023

Accepted: 2 November 2023

Published: 4 November 2023



Copyright: © 2023 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Keywords: bridge structural health monitoring; data cleaning; one-dimensional convolutional neural network; two-stage classification model

1. Introduction

In the past few decades, structural health monitoring (SHM) technology has made remarkable advancements, leading to the widespread deployment of numerous monitoring systems and, consequently, resulting in a significant surge in monitoring data [1–3]. However, the harsh operational environments, sensor malfunctions, calibration, noise, and transmission errors inevitably introduce anomalies into the data, which can significantly affect the extraction of the structural feature information. Consequently, this may lead to inaccurate assessments and conclusions and even result in severe misjudgments and accidents. Currently, the issues of data quality and anomaly identification caused by massive data growth have garnered significant attention and have become an area of extensive research [4].

As mentioned earlier, anomalies in structural health monitoring data can be categorized into two types. The first type results from structural deterioration or damage, leading to abnormal variations in monitoring data. Such anomalies are primarily identified through the exploration of indicators representing structural characteristics [5–8]. The second type arises due to data quality issues, causing monitoring data to exhibit features distinct from normal data. Anomaly detection for this category largely relies on data-driven approaches. One cost-effective solution to enhance the data-cleaning capability of existing traditional

sensors is to implement data preprocessing on raw output data for data anomaly detection instead of integrating hardware modules and analysis software, which incurs high costs for self-validation capabilities [9]. For time-invariant structures, establishing input–output mapping models is one of the most common methods used to represent data anomalies. However, in civil engineering, the influence of environmental factors and noise makes it extremely challenging to construct accurate models, rendering data-driven methods the preferred choice.

The research core of data-driven methods involves data preprocessing, feature extraction, anomaly detection, and classification. Traditional data anomaly detection relies on features such as statistical characteristics [10,11]. In recent years, with the rapid development of machine learning and deep learning technologies, these techniques have found widespread applications in civil engineering [12,13]. Notably, object detection technology has become a typical and practical solution in various scenarios, including vehicle trajectory recognition [14,15], the dynamic weighing of bridges [16], and active collision prevention for ships [17]. In the domain of data anomaly detection, high-sampling-rate acceleration data generate a large volume of complex and diverse anomaly features. Furthermore, different facilities may exhibit varying types of equipment data anomalies, making it challenging to achieve uniformity with anomaly types [18,19].

Researchers have conducted extensive studies to address the challenges posed by large datasets, diverse anomaly features, and numerous anomaly categories in data cleaning. Among various approaches, supervised learning algorithms have been widely applied. Bao et al. [18] transformed acceleration time-series data into two-dimensional images and utilized deep neural networks for data anomaly detection and classification, which were validated using monitoring data from long-span, cable-stayed bridges. Tang et al. [20] enhanced anomaly detection accuracy by incorporating the frequency domain characteristics of acceleration data alongside time series. They employed convolutional neural networks with a dual-channel input consisting of time series and spectrogram images. However, they encountered severe misclassification issues between certain categories, particularly for trend and drift types. Shajihan et al. [21] addressed misclassification problems by introducing probability density function (PDF) features. Although this improvement alleviated the misclassification of trend and drift categories, the classification performance for minor and outlier categories remained limited. Liu et al. [22] addressed data imbalance issues by utilizing unsupervised deep learning algorithms. They stacked time-series, power spectral density, and Gramian Angular Summation Field (GASF) features as a three-channel input, trained normal data using Generative Adversarial Networks (GAN), and used the output as input for convolutional neural networks to achieve data anomaly feature extraction and classification. Converting one-dimensional signals and their features into two-dimensional images to transform the data anomaly detection problem into an image feature recognition task has been a promising approach. However, due to the influence of image resolution on feature extraction and the lack of additional information in the process, employing 1D convolutional neural networks (CNNs) and statistical features for multi-class anomaly discrimination has proven to be an effective solution [19]. This method not only achieves high accuracy but also incurs lower training costs. Moreover, if the specific application scenario does not require categorizing data anomalies, the data anomaly detection task can be simplified into a binary classification problem [23–25]. Researchers [26] have effectively addressed the challenges of data imbalance and time-consuming label creation using unsupervised deep learning algorithms. By combining Gramian Angular Field (GAF) features with Generative Adversarial Networks (GANs) and autoencoders, they achieved data normality discrimination. The above studies demonstrate that utilizing general data features and dedicated models is an effective approach to enhance the generalization performance of data cleaning. Gao et al. [10] found that by extracting appropriate data features and merging labels based on these features, higher accuracy in data anomaly detection can be achieved while effectively reducing input feature dimensions, which enhances computational speed. Similarly, Samudra et al. [11] employed statistical features of monitoring data

to train random forest classifiers within decision trees, resulting in the transformation of raw data into statistical features and the subsequent reduction of input dimensions. By employing a hierarchical decision approach, an average classification accuracy of 98% was achieved. A multi-stage classification concept was introduced and validated in the anomaly classification of real bridge monitoring data [27,28]. These studies indicate that establishing multi-stage classification models for specific anomalies can effectively enhance the accuracy of classification.

Building upon these research foundations, this research establishes a two-stage framework for data anomaly classification. In the first stage, the model is based on the symmetry of acceleration signal amplitudes and frequency-domain characteristics, effectively identifying preliminary categories of data anomalies. In the second stage, a further classification is applied to an outlier category and a normal category, making the proposed framework capable of perceiving both macroscopic and detailed features. The structure of this paper is organized as follows: Section 2 introduces the data source, dataset composition, and data preprocessing methods employed. Section 3 outlines the framework and model training process of the proposed method. Section 4 presents the classification performance of the model. Section 5 concludes the entire study.

2. Data Description and Data Preprocessing

2.1. Bridge Overview and Data Set Composition

This study utilizes acceleration monitoring data from a large-span arch bridge. The bridge in question is a dual-cantilever steel box girder tied-arch bridge with a main span of 470 m and side spans of 110 m each. Fifteen acceleration sensors were installed on the bridge, as depicted in Figure 1. For the purpose of anomaly detection, acceleration monitoring data from April 2016 and December 2018 were selected, with a sampling frequency of 50 Hz. The raw data were segmented into fixed-length, time-series samples of one hour each. A total of 21,960 data sequences were obtained ($61 \text{ days} \times 24 \text{ h} \times 15 \text{ channels}$). Each data sample has a length of $180,000 \times 1$ data points. Using manual labeling, labels were assigned to each sample sequence, categorizing all samples into six distinct classes. Table 1 provides a detailed description of each class along with the corresponding sample counts. It is evident from the table that ‘minor’ anomalies and ‘missing’ data are the predominant anomaly types, whereas other anomaly types are less frequent, with noise-related anomalies only observed in April 2016. Figure 2 provides a detailed temporal distribution of anomalies for each data category.

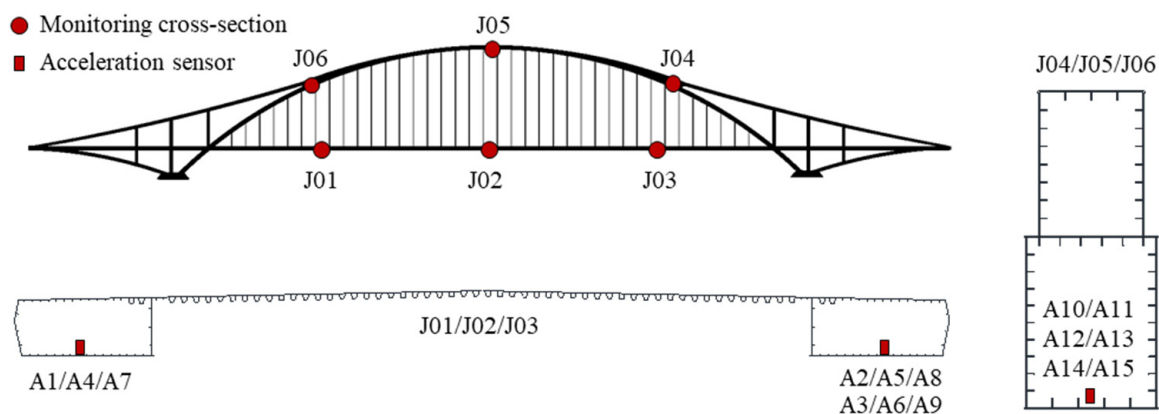


Figure 1. Placement diagram of acceleration sensors for a long-span arch bridge.

Table 1. Information about labeled data samples.

Data Patterns	Description of Data Features	Number (Ratio of Total (%))
Normal	The amplitude of a normal acceleration sequence oscillates around the vicinity of zero, and the distribution of amplitudes exhibits symmetry.	12,316 (64.50%)
Missing	When data collected or transmitted by sensors contain missing values, they are directly filled with 9999. Therefore, direct quantification is feasible without the need for model recognition.	1761 (9.22%)
Minor	Compared to a normal acceleration sequence, its amplitudes are reduced by 1–2 orders of magnitude. Additionally, influenced by factors such as data storage precision, the occurrence of square wave phenomena is more likely.	3956 (20.72%)
Biased	The acceleration exhibits two distinct forms within the ranges of greater than 0 and less than 0, disrupting the symmetric distribution characteristic of acceleration amplitudes.	363 (1.90%)
Outlier	The acceleration time series exhibits significant fluctuations in amplitude.	123 (0.64%)
Noise	Due to the influence of noise, the acceleration is submerged within high-amplitude noise, with no apparent fluctuations.	576 (3.02%)

**Figure 2.** Distribution of labeled data samples.

2.2. Data Preprocessing

2.2.1. Missing Category

In this study, the raw data from acceleration sensors installed on the bridge are subject to instances of data loss in which the affected sampling points are directly assigned a value of 9999. Consequently, a quantitative approach is employed to assess the occurrence and extent of missing data within the segment. The labeling of manually identified missing data anomalies is based on a criterion wherein data samples with a missing data proportion exceeding 40% are categorized as such. This specific anomaly category, amenable to direct quantification, is purposefully excluded from subsequent tasks involving anomaly detection. Consequently, the training dataset is purged of any influence stemming from this category.

2.2.2. Features for Classification

Addressing the issue of anomaly classification caused by data quality requires feature extraction to capture the fundamental distinctions between normal data and data affected by data quality issues. Acceleration data, characterized by high sampling frequencies

and substantial data volumes, present challenges in terms of human-intensive manual assessment for anomaly detection. In this research, a comprehensive analysis of vibration acceleration data from bridge structures reveals a distinctive characteristic of normal acceleration: its amplitude exhibits symmetry around zero and conforms to a normal distribution pattern. This pattern represents the essential disparity between normal and anomalous data.

Inversed Probability Density Function (IPDF)

Jian et al. [19] introduced the relative frequency distribution histogram (RFDH) metric to characterize anomalies in acceleration data. This feature utilizes the symmetric distribution characteristic of acceleration data around the zero-amplitude point and incorporates various statistical features, such as kurtosis and skewness. However, it is important to note that its identification performance is suboptimal when dealing with outlier patterns.

In this paper, the same feature is employed with only slight modifications. Firstly, assuming the distribution range of acceleration amplitudes is within ± 50 mg, this range is divided into 1025 small intervals. The number of samples in each interval is counted, and the values are normalized by the maximum count. Subsequently, 1 is subtracted from the values of non-zero intervals, excluding the maximum value from this calculation. This process culminates in the formation of the IPDF feature.

Figure 3 depicts the IPDF features of various data patterns. It is evident that the normal data type exhibits a single-peaked normal distribution centered around 0. The minor pattern's amplitudes concentrate in the vicinity of 0. The outlier type, due to the substantial disturbances from anomalies with large amplitudes, results in a broader amplitude distribution. However, the distribution near 0 still demonstrates the characteristics of a single-peaked normal distribution. The biased category exhibits asymmetric distribution in the IPDF metric due to the presence of abnormally low acceleration amplitudes in the negative range. The noise class manifests as a bimodal distribution in the amplitude domain.

Power Spectral Density (PSD)

To compensate for the deficiencies in amplitude domain features, this paper introduces frequency-domain features by calculating the power spectral density of the acceleration signal. Segments of 2048 sampling points are used as segments, with an overlap length of 1024 to calculate the power spectral density, and the amplitude of the power spectrum is normalized using the maximum value. Figure 3 illustrates the power spectral features of various data patterns. In the normal category, peaks of varying heights appear in the 0–10 Hz range, reflecting the inherent dynamic characteristics of the structure. The minor category exhibits two abnormal modes: One involves abnormal amplitudes but normal vibrations, resulting in relatively normal peak values in the power spectral density due to lower vibrational energy. However, some high-frequency noise effects are accentuated due to the smaller vibrational energy. The other mode occurs when data amplitudes fluctuate near the data storage precision, leading to square wave characteristics in the time domain and completely abnormal spectral characteristics. For outlier patterns, high-energy disturbances lead to energy concentration in the power spectrum around 0. Although the biased class shows peaks in the 0–10 Hz range, these frequencies do not reflect the inherent dynamic characteristics of the structure when compared to normal acceleration data. Lastly, the noise type, due to high-energy noise interference, predominantly showcases noise frequency components within its spectrum, completely submerging the normal frequency components of the signal.

Ratio of Peak-to-Valley (RPV)

IPDF and PSD extract features from the amplitude and frequency domains of the raw acceleration data, respectively. However, it is important to note that these types of features emphasize the macroscopic properties of the data, reflecting overall characteristics such as acceleration amplitude symmetry and frequency components. For local data anomalies,

such as local large-amplitude jumps, these features often struggle to meet application requirements. In this regard, this section proposes the RPV indicator to address local anomalies in the acceleration data.

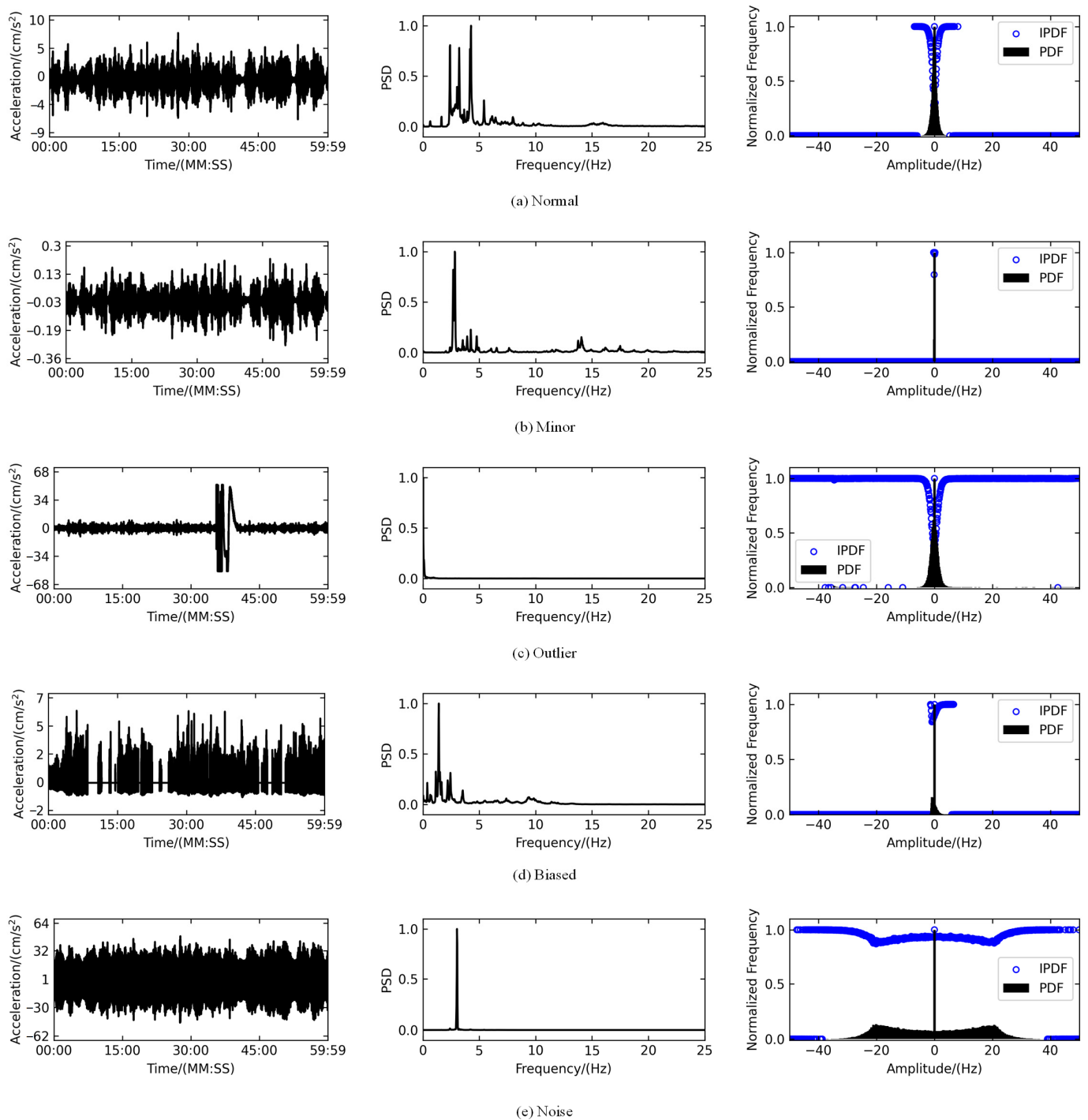


Figure 3. Probability density function and power spectral density of data patterns.

The calculation of the RPV indicator involves picking the maximum peak and minimum valley values of the original time series within equally spaced sliding windows and also within the peak and valley values of the time series within a constrained interval, as defined in Equation (1).

$$RPV = \frac{|p_r + v_r|}{|p_c - v_c|} \quad (1)$$

where p_r represents the maximum peak value within the segment of the raw acceleration data, v_r represents the minimum valley value within the segment of the raw acceleration data, p_c represents the maximum peak value within the segment of the constrained acceleration data, and v_c represents the minimum valley value within the segment of the constrained acceleration data.

The calculation of peak and valley values for this indicator is carried out under two conditions: one using the original acceleration data and the other using the constrained acceleration data, as illustrated in Figure 4. In this study, segment lengths of 1000 sample points are utilized, with a step size of 500 sample points. Based on the normal distribution characteristics of acceleration amplitudes, it is assumed that within the specified segment length range, normal acceleration amplitudes should be distributed within the range of ± 3 standard deviations from the mean. Values exceeding this range can be approximated as anomalous jump points. It should be noted here that to mitigate the influence of extremely small values within local segment data, the standard deviation is selected as the greater value between 0.5 and the computed standard deviation of the current segment.

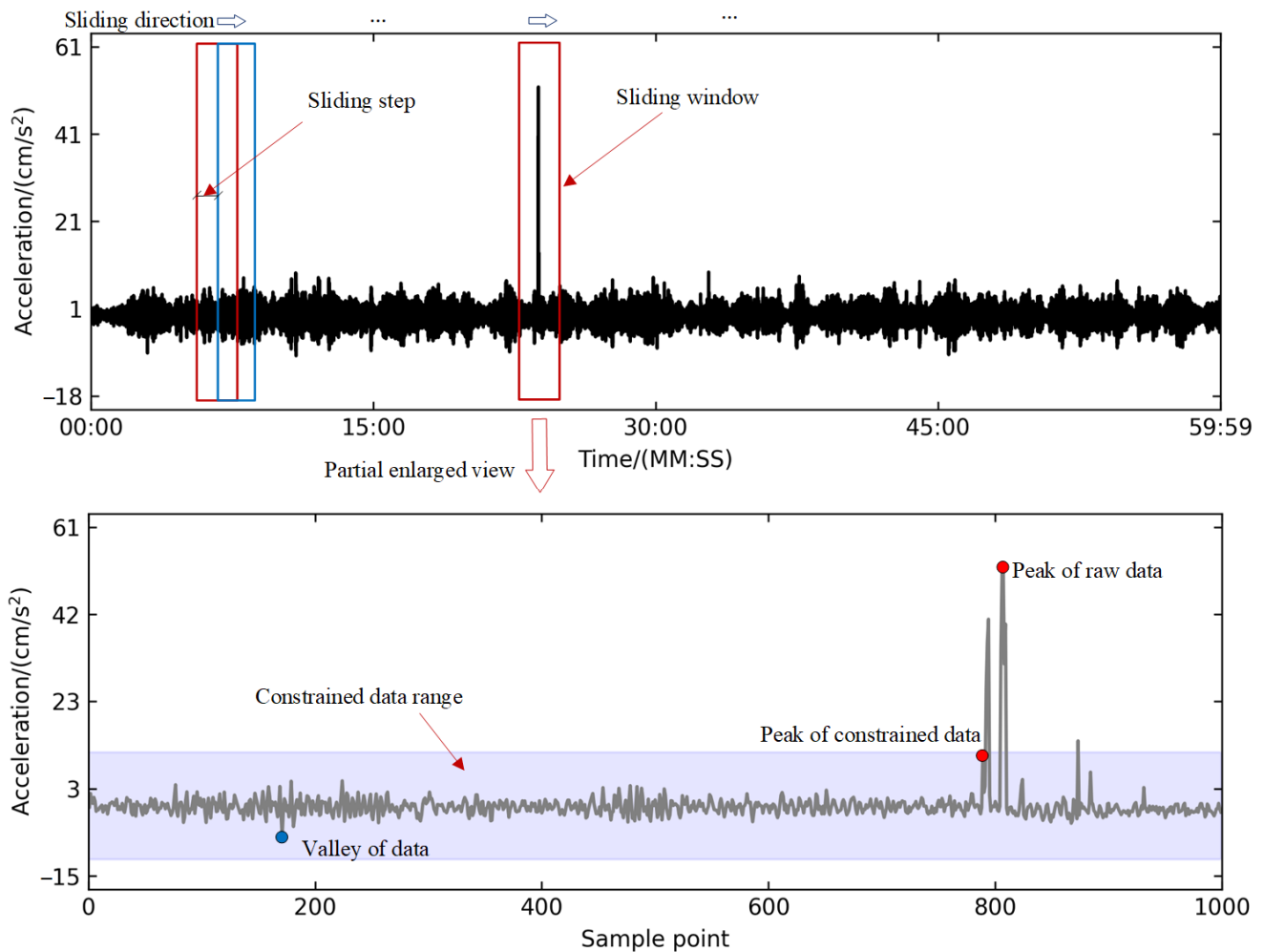


Figure 4. Schematic diagram of the RVP calculation process.

After calculating the indicator using Equation (1), a comparison is made between the data of the normal category and the outlier category. As illustrated in Figure 5, it is evident that the RPV of the normal category data does not exceed 1, while the outlier category exhibits points exceeding 1. This provides a straightforward way to visually identify abnormal patterns in the two categories of data.

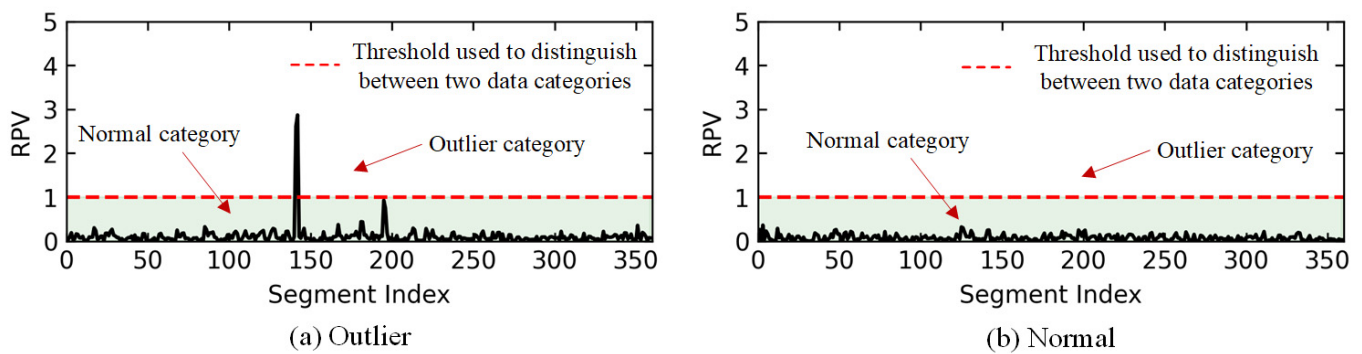


Figure 5. Comparison of RPV between normal and outlier categories of samples.

3. Proposed Framework and Model Training

3.1. Framework of the Proposed Method

This paper proposes a two-stage model for data anomaly detection, as illustrated in Figure 6. Firstly, the PSD and IPDF features are extracted from the raw accelerometer time series. These features are combined into a dual-channel input and used to train a one-dimensional convolutional neural network. The trained model is then used for classifying five types of data anomalies. As mentioned earlier, the focus of these two types of input features is on the global characteristics of the data, making them less sensitive to local data anomalies. Building upon this, the paper further differentiates between the most confusing outlier and normal categories. Data classified into these two categories are subjected to the extraction of the RPV metric. A threshold-based filtering process is applied to ascertain the presence of confusion between these two categories, thereby enhancing the sensitivity of the classification model toward local data anomalies.

In this study, a one-dimensional convolutional neural network (1D CNN) is employed as the feature extraction and classification model. CNNs possess characteristics such as sparse representation, parameter sharing, and equivariant representation that contribute to high computational efficiency. The architecture of the CNN used in this study is illustrated in Table 2.

Table 2. Architecture of the CNN.

Network Layer	Parameter Setting
Input layer	Size: 256 (Batch size) $\times$ 2 (Channels) $\times$ 1025 (Length of features) Number of channels in each input feature: 2 Number of channels in each output feature: 16
Convolutional layer 1	Kernel size: 3 $\times$ 1 Stride: 1 Padding: same Batch Normalization: True Activation function: ReLU
Pooling layer 1	Kernel size: 2 $\times$ 1 Stride: 2 Number of channels in each input feature: 16 Number of channels in each output feature: 32
Convolutional layer 2	Kernel size: 7 $\times$ 1 Stride: 1 Padding: same Batch Normalization: True Activation function: ReLU
Pooling layer 2	Kernel size: 2 $\times$ 1 Stride: 2

Table 2. Cont.

Network Layer	Parameter Setting
Convolutional layer 3	Number of channels in each input feature: 32 Number of channels in each output feature: 64 Kernel size: 3×1 Stride: 1 Padding: same Batch Normalization: True Activation function: ReLU
Pooling layer 3	Kernel size: 2×1 Stride: 2
Flatten layer	Size: 256×8128
Fully connected layer	Size of each input feature: 8128 Size of each output feature: 100 Activation function: ReLU
Fully connected layer	Size of each input feature: 100 Size of each output feature: 5

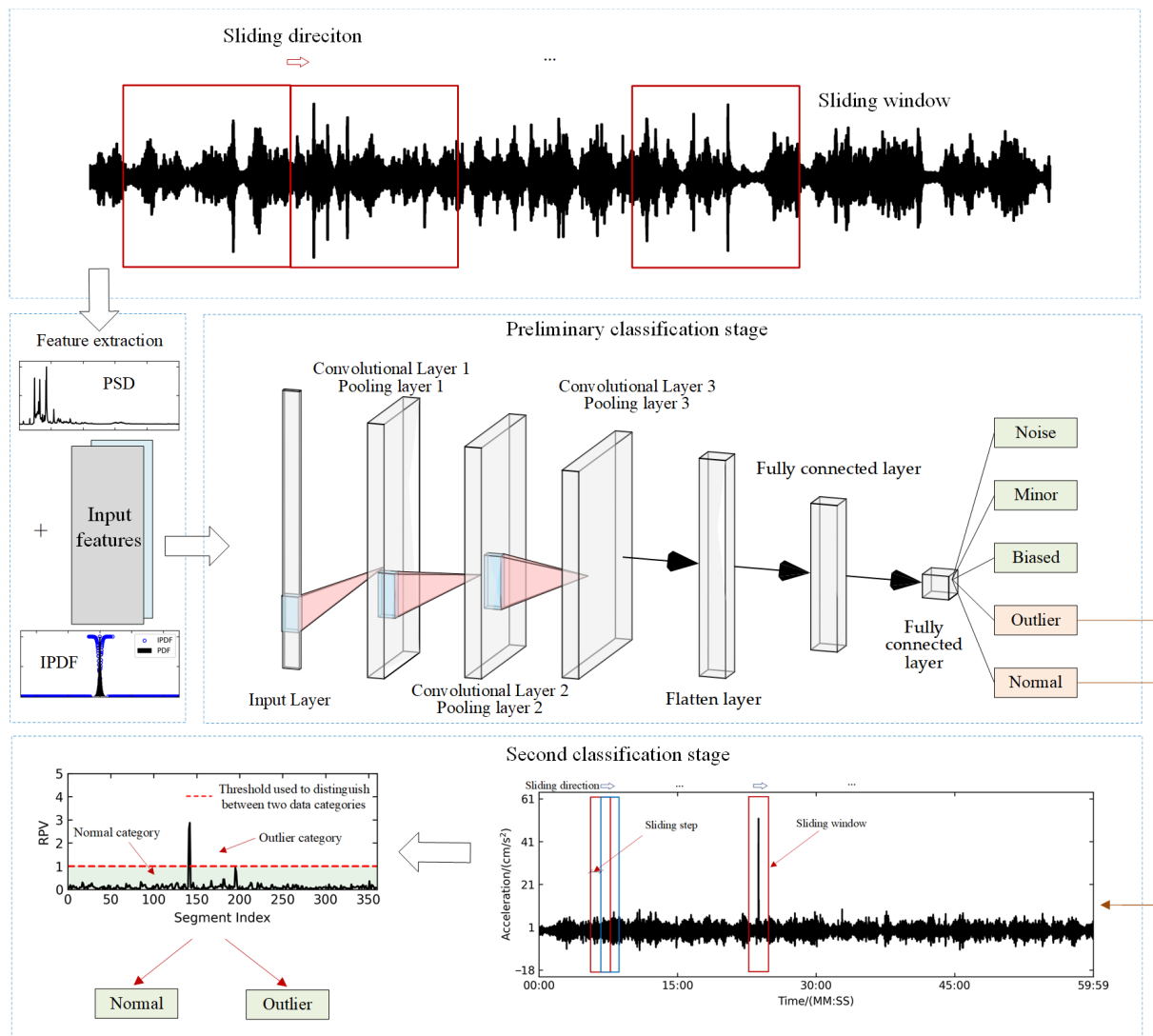


Figure 6. Framework of the proposed method.

3.2. Performance Evaluation

As shown in Equations (2)–(5), accuracy, precision, recall, and F1 score are commonly used to measure the classification performance of models and are suitable for both binary and multi-class problems. Accuracy measures the proportion of correctly classified samples to the total number of samples. Precision is the ratio of correctly predicted positive samples to the total number of samples predicted as positive. Recall is the ratio of correctly predicted positive samples to the total number of true positive samples. The F1 score combines precision and recall and is calculated as the harmonic mean of the two.

$$accuracy = \frac{TP + TN}{TP + FN + FP + TN} \quad (2)$$

$$precision = \frac{TP}{TP + FP} \quad (3)$$

$$recall = \frac{TP}{TP + FN} \quad (4)$$

$$F_1 = 2 \times \frac{precision \times recall}{precision + recall} \times 100\% \quad (5)$$

where TP represents the number of samples that are truly positive and are correctly predicted as positive by the model. TN represents the number of samples that are truly negative and are correctly predicted as negative by the model. FP represents the number of samples that are actually negative but are incorrectly predicted as positive by the model. FN represents the number of samples that are actually positive but are incorrectly predicted as negative by the model. For multi-class problems, you can treat the current class as the positive class and group all other classes as the negative class, making these metrics equally applicable.

3.3. Model Training

This study employs 70% of the labeled samples for each category as the training set, 10% as the validation set, and 20% as the test set. The training process spans 100 epochs, with a batch size of 256. The cross-entropy loss function is utilized, coupled with the Adam optimization algorithm. The initial learning rate is set at 0.001, and weight decay is incorporated to mitigate the overfitting of the model.

The accuracy of the training and validation sets during training is shown in Figure 7. Due to the presence of an imbalanced dataset, the training accuracy quickly reaches 99%. After around 45 epochs, the training accuracy stabilizes, and there is little significant improvement in the validation set accuracy.

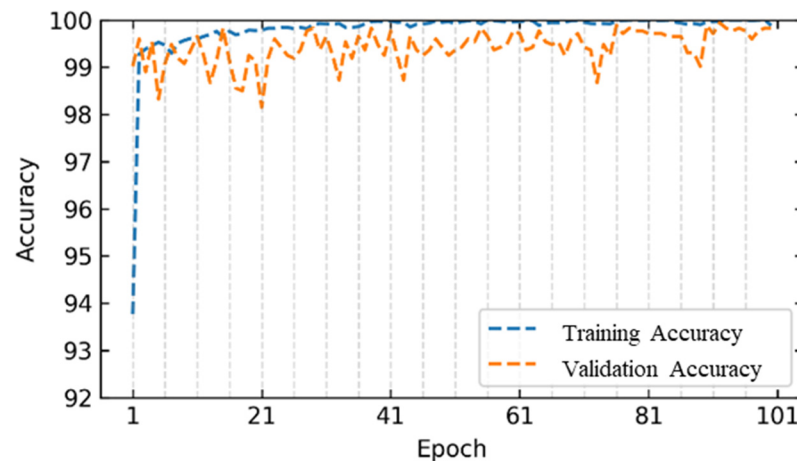


Figure 7. The accuracy of the training set and the validation set.

4. Results

The trained model is applied to the testing dataset for classification tasks. A comparison is made to assess the impact of adding a second-stage classification model on the final classification results. The confusion matrices under the two scenarios are shown in Figure 8. It can be observed that after the second-stage classification, the confusion issue between the normal and outlier categories has improved. This framework not only resolves the confusion problem between the two abnormal data types but also extends the perceptual ability of the 1D CNN classification model to details, with clear principles and high result credibility.

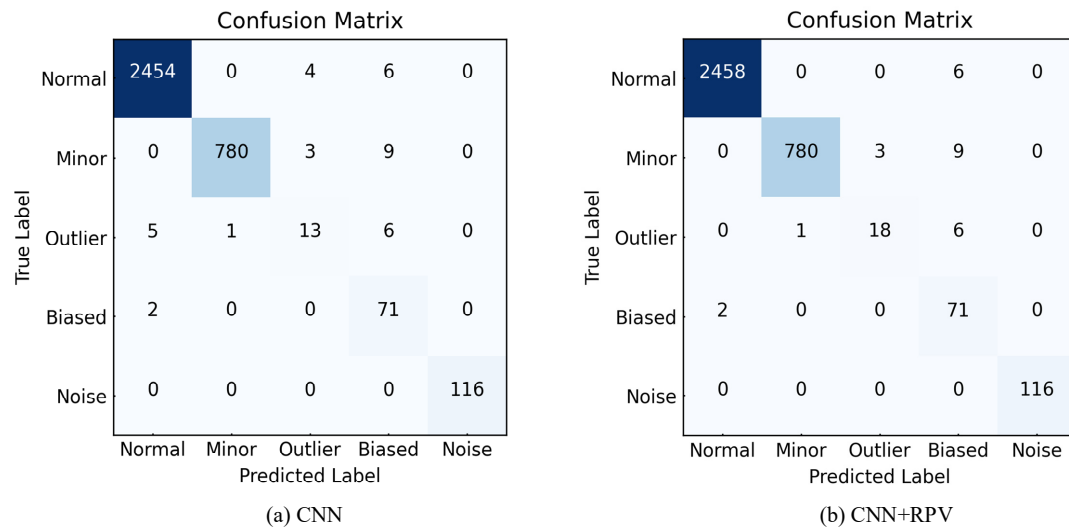


Figure 8. The confusion matrix of a CNN model only and a two-stage classification model.

For the other types, after passing them through the 1D CNN model under the current bridge's data anomaly characteristics, they already exhibited high classification accuracy. There are a few misclassifications between the minor and outlier categories, mainly due to the coupling phenomenon of multiple abnormal types within data sample segments. This is influenced by subjectivity during manual labeling. In practice, as long as the normal category obtained from the classification has high credibility, it meets the requirements for data analysis. The exact classification of data anomalies into specific categories is not the primary concern.

The performance evaluation metrics mentioned earlier were used to further explain the results for the two-stage approach, and the results are presented in Table 3. It can be observed that the calculations for normal, minor, and noise are all close to 100%, indicating a good classification performance. However, the results for the outlier and biased categories show a slightly lower performance. This discrepancy can be attributed to two main factors:

Table 3. Performance evaluation metrics of the model on the test set.

Performance Evaluation Metrics	Normal	Minor	Outlier	Biased	Noise
Precision	99.92%	99.87%	85.71%	77.17%	100%
Recall	99.76%	98.48%	72.00%	97.26%	100%
F_1	99.84%	99.17%	78.26%	86.06%	100%
Accuracy			99.22%		

Class Imbalance: The severe class imbalance in the dataset leads to insufficient training of the model. The presence of a majority class can dominate the learning process, resulting in less emphasis on the minority classes, like outlier and biased.

Subjectivity in Labeling: The coupling of multiple types of data anomalies during manual labeling introduces subjectivity. This subjectivity might lead to ambiguous cases that are difficult to classify, particularly for the outlier and biased categories.

Moreover, the limited size of the test set amplifies the impact of misclassifications, even if the number of misclassified samples is minimal. Therefore, while the overall classification performance seems satisfactory, the influence of misclassifications on the metrics is significant due to the small test sample size.

5. Conclusions

To address the challenge of identifying anomalous data within the vast accumulation of acceleration data in bridge health monitoring systems, this study introduces a two-stage data cleaning method. The proposed approach initially transforms the raw data sequences into IPDF and PSD dual-channel features, employing a one-dimensional convolutional neural network for multi-class classification tasks. Moreover, the RPV metric is incorporated to address misclassification issues between the normal and outlier categories. The results from the two-stage classification model exhibit an overall accuracy exceeding 99%, effectively ameliorating misclassification problems associated with outlier and normal categories. This augmentation enhances the model's ability to perceive local disturbances with substantial magnitudes.

However, due to sample imbalance and the coupling of various data anomalies, there remains room for enhancement in the model's classification performance. While the determination credibility for the normal category has been considerably improved and misclassification between outlier and normal categories has been effectively mitigated, attention must still be given to misclassification issues among other data categories. In future research, focus should be directed toward constructing multi-label classification models and addressing the problem of anomaly localization within the data.

Author Contributions: Y.X. carried out the studies, participated in data collection and analysis, and drafted the manuscript. Y.Z. conceived of the study, participated in its design and coordination, and helped to draft the manuscript. J.Z. participated in the design of the study as well as the writing—review and editing. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: All data, models, and codes to support the findings of this study are available from the corresponding author upon request.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Sun, L.; Shang, Z.; Xia, Y.; Bhowmick, S.; Nagarajaiah, S. Review of bridge structural health monitoring aided by big data and artificial intelligence: From condition assessment to damage detection. *J. Struct. Eng.* **2020**, *146*, 04020073. [\[CrossRef\]](#)
2. Ou, J.; Li, H. Structural health monitoring in mainland China: Review and future trends. *Struct. Health Monit.* **2010**, *9*, 219–231.
3. Abdulkarem, M.; Samsudin, K.; Rokhani, F.; Rasid, M.F.A. Wireless sensor network for structural health monitoring: A contemporary review of technologies, challenges, and future direction. *Struct. Health Monit.* **2020**, *19*, 693–735. [\[CrossRef\]](#)
4. Feng, D.; Feng, M.; Mao, X.; Zhu, H. A review of data processing methods for bridge structural health monitoring. *Struct. Control Health Monit.* **2020**, *27*, e2518.
5. Avci, O.; Abdeljaber, O.; Kiranyaz, S.; Hussein, M.; Gabbouj, M.; Inman, D.J. A review of vibration-based damage detection in civil structures: From traditional methods to Machine Learning and Deep Learning applications. *Mech. Syst. Signal Process.* **2021**, *147*, 107077. [\[CrossRef\]](#)
6. An, Y.; Chatzi, E.; Sim, S.; Laflamme, S.; Blachowski, B.; Ou, J. Recent progress and future trends on damage identification methods for bridge structures. *Struct. Control Health Monit.* **2019**, *26*, e2416. [\[CrossRef\]](#)
7. Daneshvar, M.; Saffarian, M.; Jahangir, H.; Sarmadi, H. Damage identification of structural systems by modal strain energy and an optimization-based iterative regularization method. *Eng. Comput.* **2023**, *39*, 2067–2087. [\[CrossRef\]](#)

8. Qu, Y.; Hu, N.; Li, Z. Anomaly detection in structural health monitoring using variational Bayesian learning. *Struct. Health Monit.* **2015**, *14*, 359–368.
9. Feng, Z.; Wang, Q.; Shida, K. A review of self-validating sensor technology. *Sens. Rev.* **2007**, *27*, 48–56. [[CrossRef](#)]
10. Gao, K.; Chen, Z.; Weng, S.; Zhu, H.P.; Wu, L.Y. Detection of multi-type data anomaly for structural health monitoring using pattern recognition neural network. *Smart Struct. Syst.* **2022**, *29*, 129–140.
11. Samudra, S.; Barbosh, M.; Sadhu, A. Machine learning-assisted improved anomaly detection for structural health monitoring. *Sensors* **2023**, *23*, 3365. [[CrossRef](#)] [[PubMed](#)]
12. Bao, Y.; Li, H. Machine learning paradigm for structural health monitoring. *Struct. Health Monit.* **2021**, *20*, 1353–1372. [[CrossRef](#)]
13. Teng, S.; Chen, G.; Liu, Z.; Cheng, L.; Sun, X. Multi-sensor and decision-level fusion-based structural damage detection using a one-dimensional convolutional neural network. *Sensors* **2021**, *21*, 3950. [[CrossRef](#)]
14. Ge, L.; Dan, D.; Li, H. An accurate and robust monitoring method of full-bridge traffic load distribution based on YOLO-v3 machine vision. *Struct. Control Health Monit.* **2020**, *27*, e2636. [[CrossRef](#)]
15. Xia, Y.; Jian, X.; Yan, B.; Su, D. Infrastructure safety oriented traffic load monitoring using multi-sensor and single camera for short and medium span bridges. *Remote Sens.* **2019**, *11*, 2651. [[CrossRef](#)]
16. Dan, D.; Ying, Y.; Ge, L. Digital twin system of bridges group based on machine vision fusion monitoring of bridge traffic load. *IEEE T. Intell. Transp.* **2021**, *23*, 22190–22205. [[CrossRef](#)]
17. Tang, C.; Chen, M.; Zhao, J.; Liu, T.; Liu, K.; Yan, H.; Xiao, Y. A novel ship trajectory clustering method for finding overall and local features of ship trajectories. *Ocean Eng.* **2021**, *241*, 110108. [[CrossRef](#)]
18. Bao, Y.; Tang, Z.; Li, H.; Zhang, Y. Computer vision and deep learning-based data anomaly detection method for structural health monitoring. *Struct. Health Monit.* **2019**, *18*, 401–421. [[CrossRef](#)]
19. Jian, X.; Zhong, H.; Xia, Y.; Sun, L. Faulty data detection and classification for bridge structural health monitoring via statistical and deep-learning approach. *Struct. Control Health Monit.* **2021**, *28*, e2824. [[CrossRef](#)]
20. Tang, Z.; Chen, Z.; Bao, Y.; Li, H. Convolutional neural network-based data anomaly detection method using multiple information for structural health monitoring. *Struct. Control Health Monit.* **2019**, *26*, e2296. [[CrossRef](#)]
21. Shajihan, S.; Wang, S.; Zhai, G.; Spencer, B.F., Jr. CNN based data anomaly detection using multi-channel imagery for structural health monitoring. *Smart Struct. Syst.* **2022**, *29*, 181–193.
22. Liu, G.; Niu, Y.; Zhao, W.; Duan, Y.; Shu, J. Data anomaly detection for structural health monitoring using a combination network of GANomaly and CNN. *Smart Struct. Syst.* **2022**, *29*, 53–62.
23. Ni, F.; Zhang, J.; Noori, M. Deep learning for data anomaly detection and data compression of a long-span suspension bridge. *Comput.-Aided Civ. Inf.* **2020**, *35*, 685–700. [[CrossRef](#)]
24. Zhang, Y.; Lei, Y. Data anomaly detection of bridge structures using convolutional neural network based on structural vibration signals. *Symmetry* **2021**, *13*, 1186. [[CrossRef](#)]
25. Sony, S.; Gamage, S.; Sadhu, A.; Samarabandu, J. Multiclass damage identification in a full-scale bridge using optimally tuned one-dimensional convolutional neural network. *J. Comput. Civil Eng.* **2022**, *36*, 04021035. [[CrossRef](#)]
26. Mao, J.; Wang, H.; Spencer, J. Toward data anomaly detection for automated structural health monitoring: Exploiting generative adversarial nets and autoencoders. *Struct. Health Monit.* **2021**, *20*, 1609–1626. [[CrossRef](#)]
27. Zhang, H.; Lin, J.; Hua, J.; Gao, F.; Tong, T. Data anomaly detection for bridge SHM based on CNN combined with statistic features. *J. Nondestruct. Eval.* **2022**, *41*, 28. [[CrossRef](#)]
28. Li, S.; Jin, L.; Qiu, Y.; Zhang, M.; Wang, J. Signal anomaly detection of bridge SHM system based on two-stage deep convolutional neural networks. *Struct. Eng. Int.* **2023**, *33*, 74–83. [[CrossRef](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.